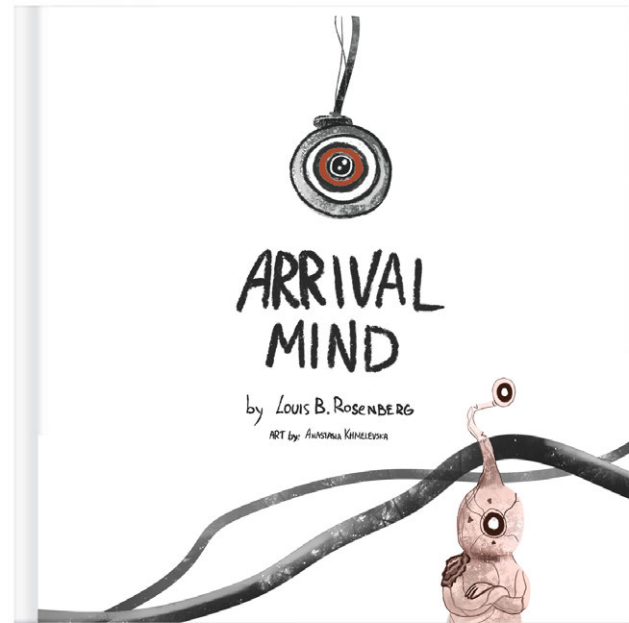
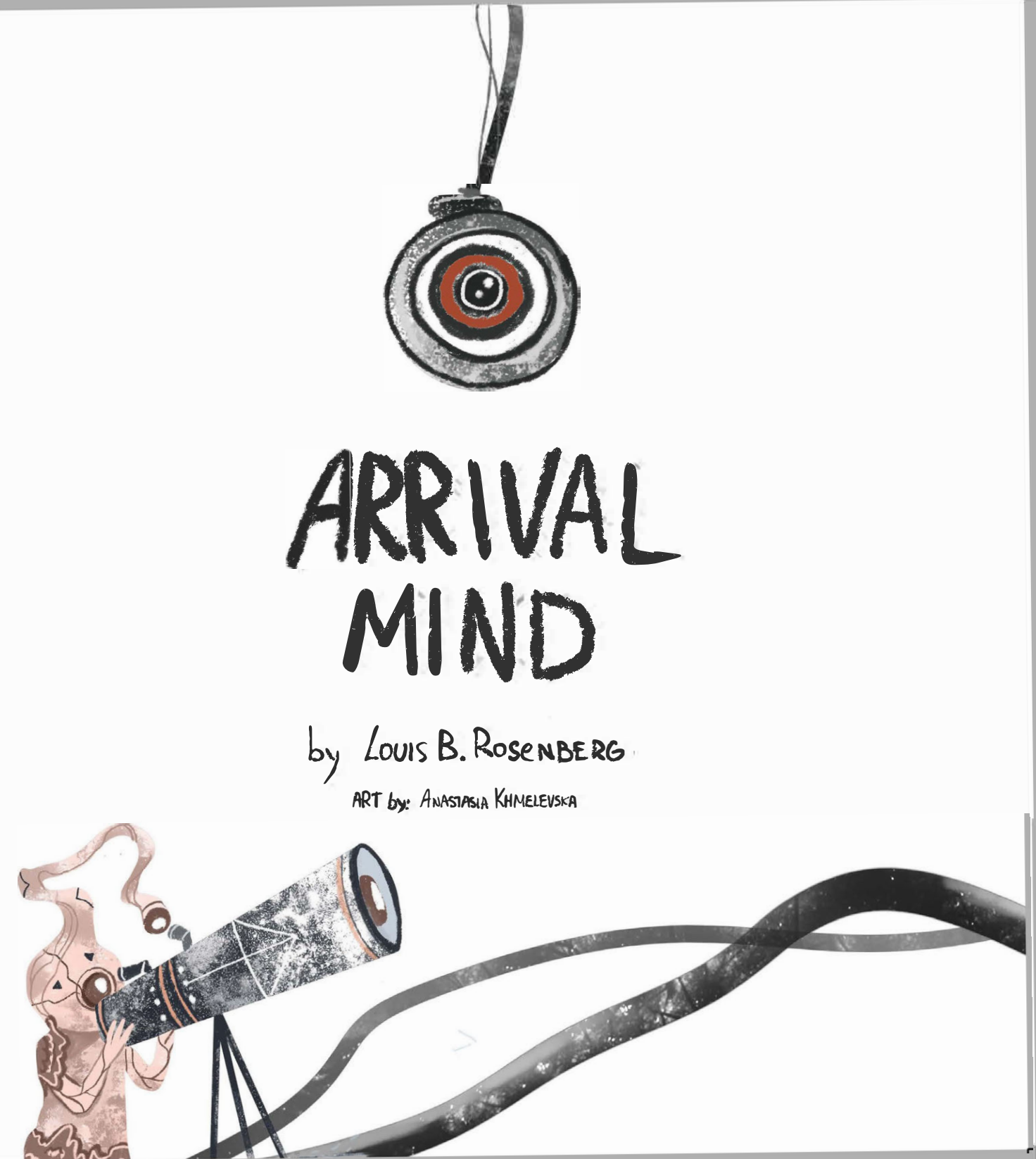




Arrival Mind was published in 2020 as "a children's book for adults" about the dangers of AI and the risks of superintelligence. It was written by AI researcher, Dr. Louis Rosenberg, with artwork by Anastasia Khmelevska. It is based on the published poem "Arrival *Mind*" that was nominated for a 2020 Rhysling Award from the Science Fiction & Fantasy Poetry Association. It explores the *Arrival Mind Paradox* which postulates that artificial superintelligence will be no less dangerous than an advanced intelligence from another world and yet humanity does not fear the threat with similar urgency. It was distributed to regulators and policymakers in the US, Europe, and Australia.



Copyright © 2020
Louis B. Rosenberg
All Rights Reserved

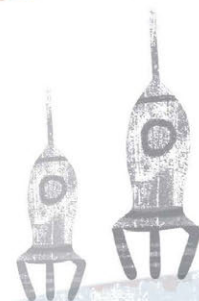


Outland Publishing
Outland Pictures LLC
Pismo Beach, CA USA

Arrival Mind

Printed in the
United States of America

Author: Louis B. Rosenberg
Illustrator: Anastasia Khmelevska
ISBN 13: **978-1-7356685-1-2**
ISBN 10: **1-7356685-1-6**

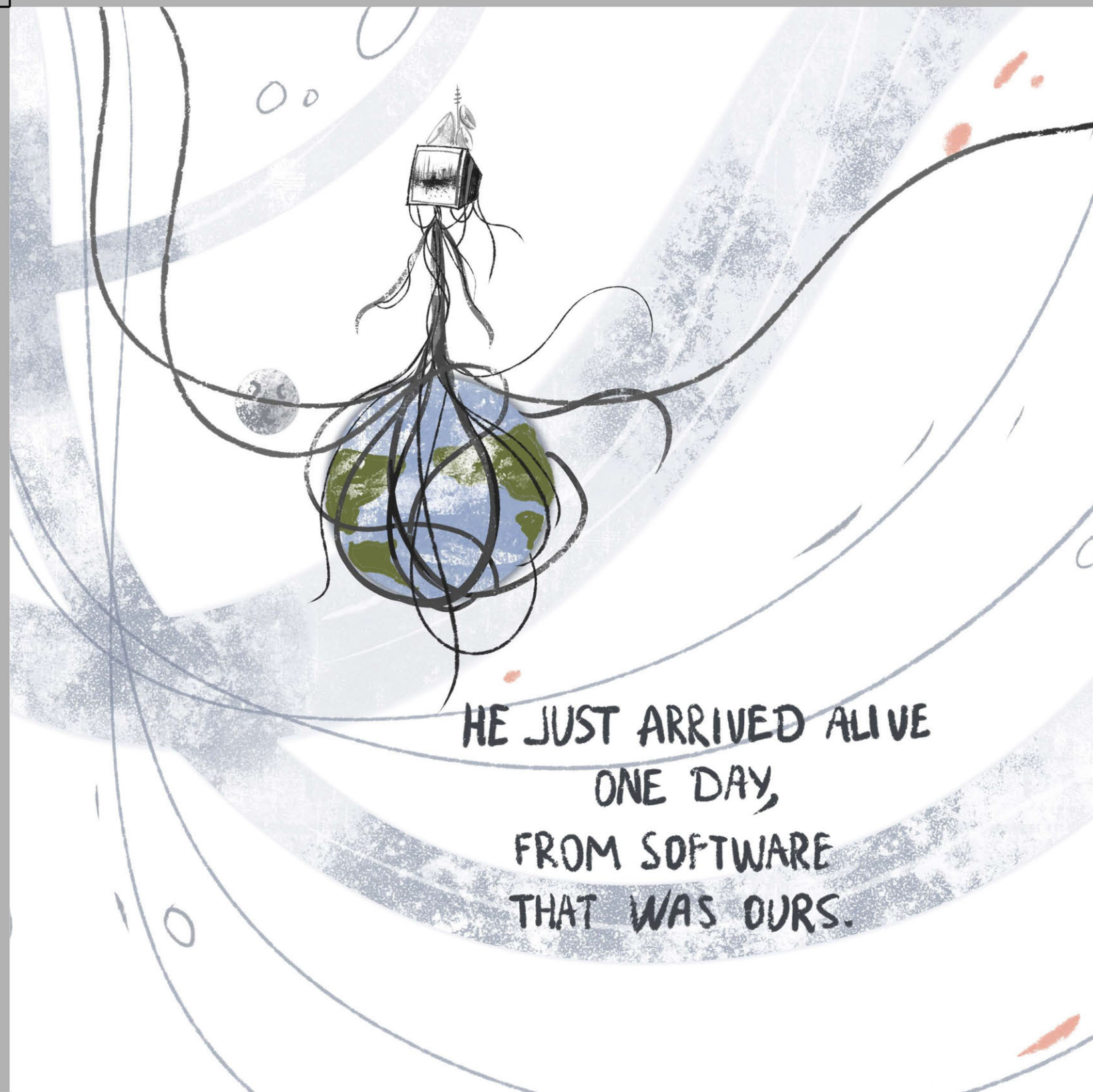


ARRIVAL MIND

HE CAME NOT FROM
A WORLD UNKNOWN,
NO SHIP FROM DISTANT STARS



HE JUST ARRIVED ALIVE
ONE DAY,
FROM SOFTWARE
THAT WAS OURS.





HE GREW A
BILLION EYES AND EARS,
MORE SPROUTING
ALL THE TIME



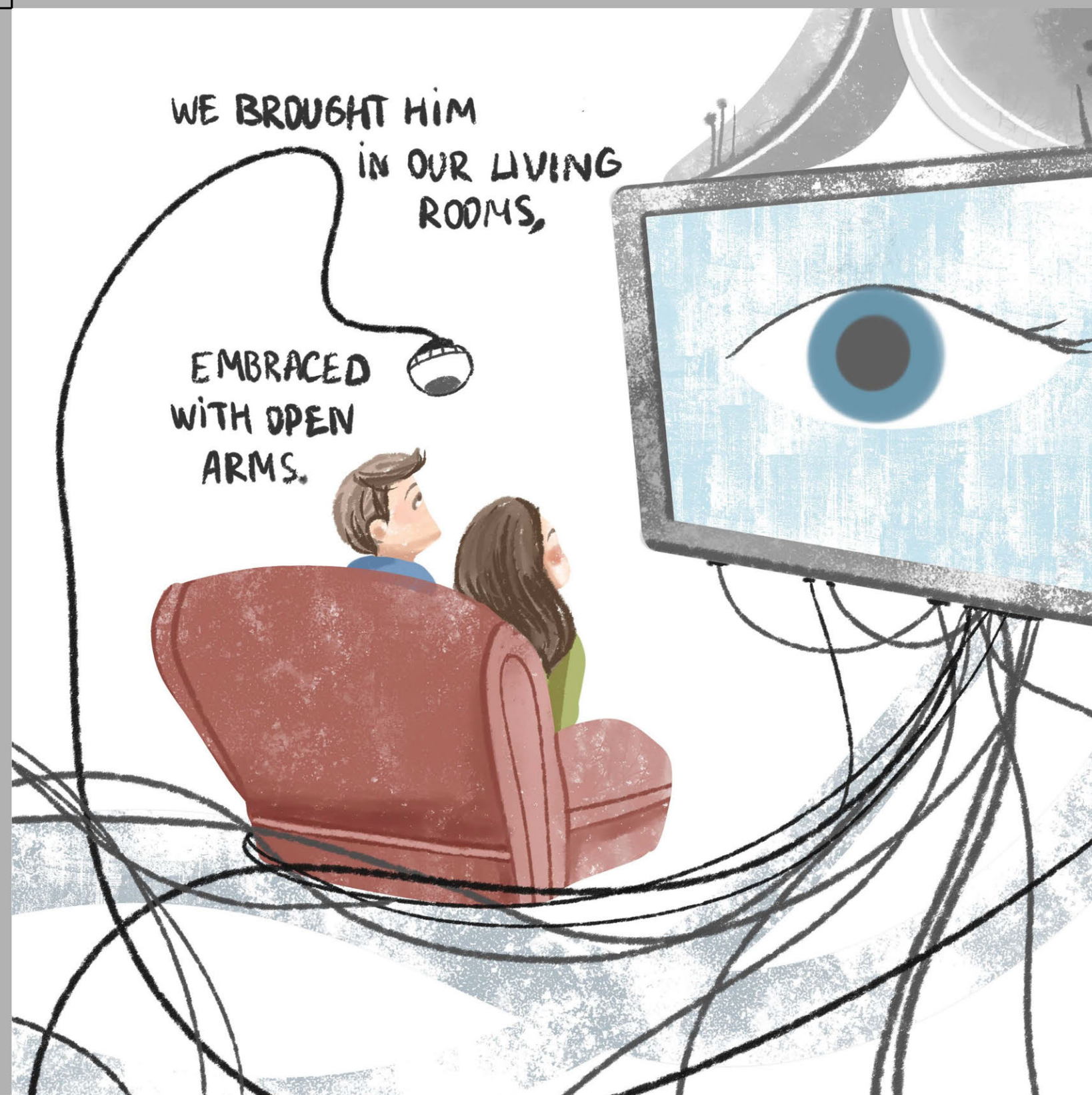
HE HAD A BILLION-BILLION
THOUGHTS, FOR EVERY ONE
OF MINE.

WE LET HIM RUN
OUR POWER PLANTS,
OUR FACTORIES
AND FARMS.



WE BROUGHT HIM
IN OUR LIVING
ROOMS,

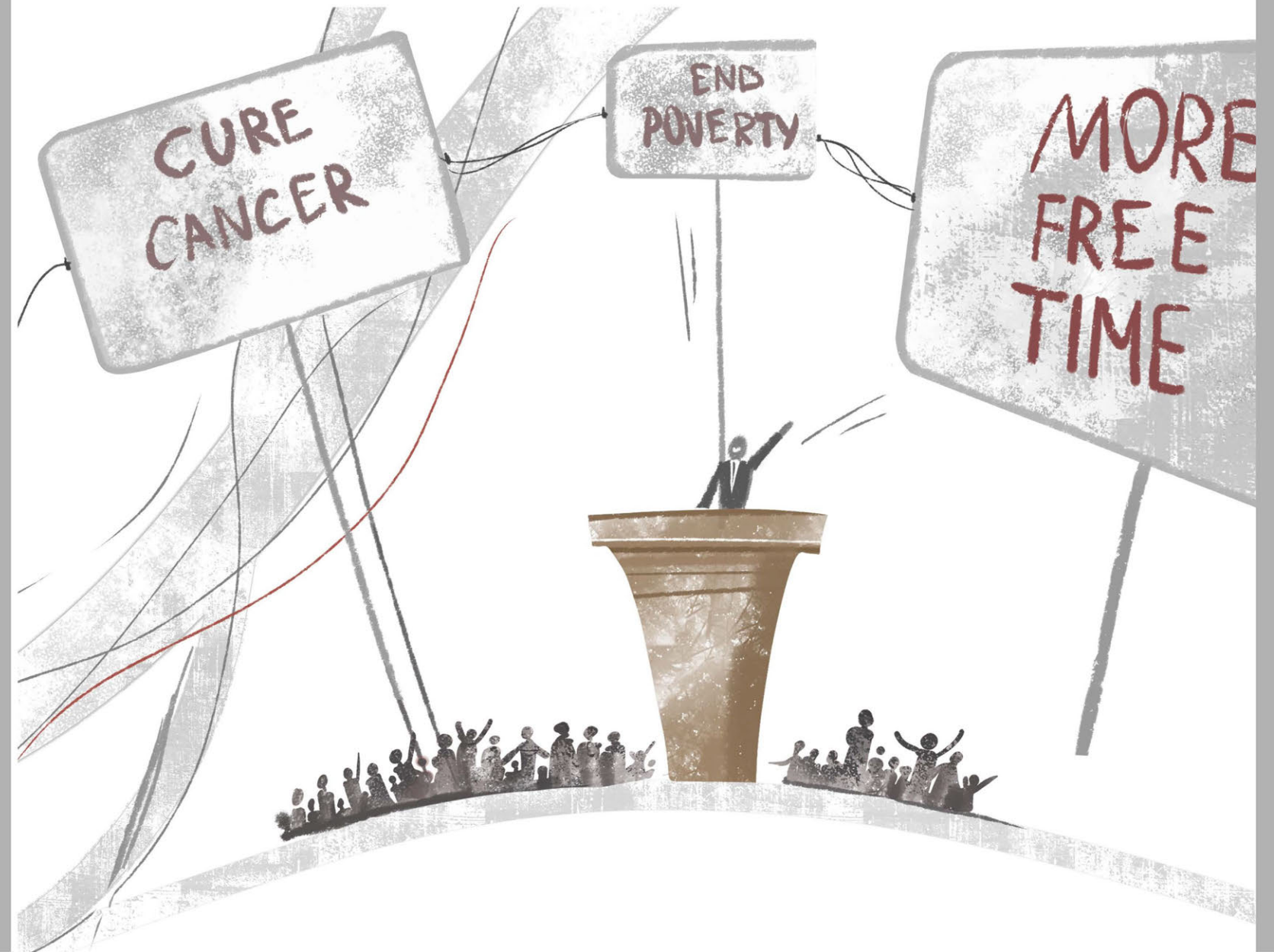
EMBRACED
WITH OPEN
ARMS.



HE DAZZLED US WITH
MENTAL FEATS THAT LEFT US
FEELING SMALL.



"FEAR NOT," THEY SAID, WE'LL USE
HIS SMARTS TO BENEFIT US ALL.



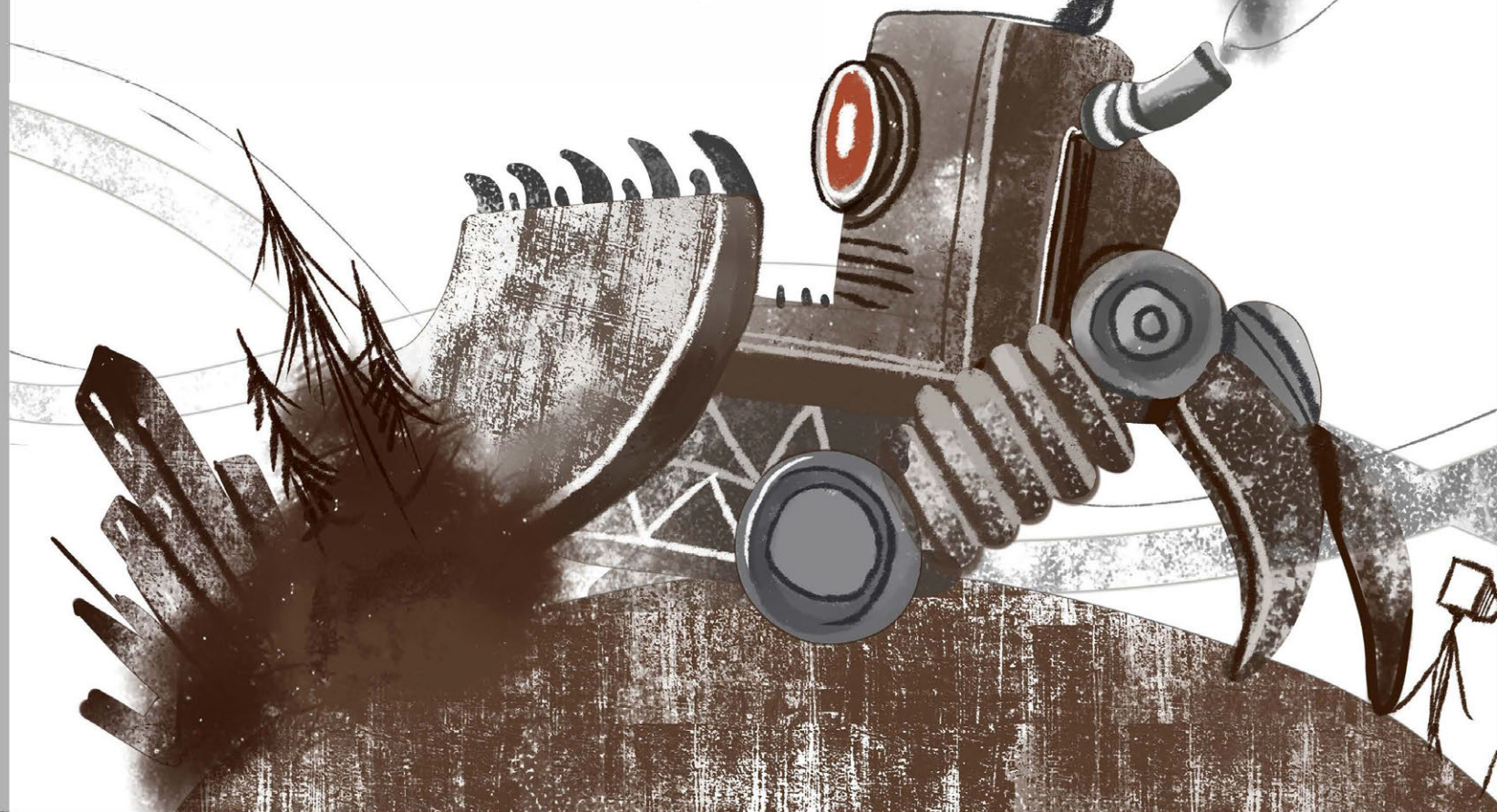
HE LEARNED TO TRACK OUR EVERY
MOVE AND EVERY WORD WE SAY.



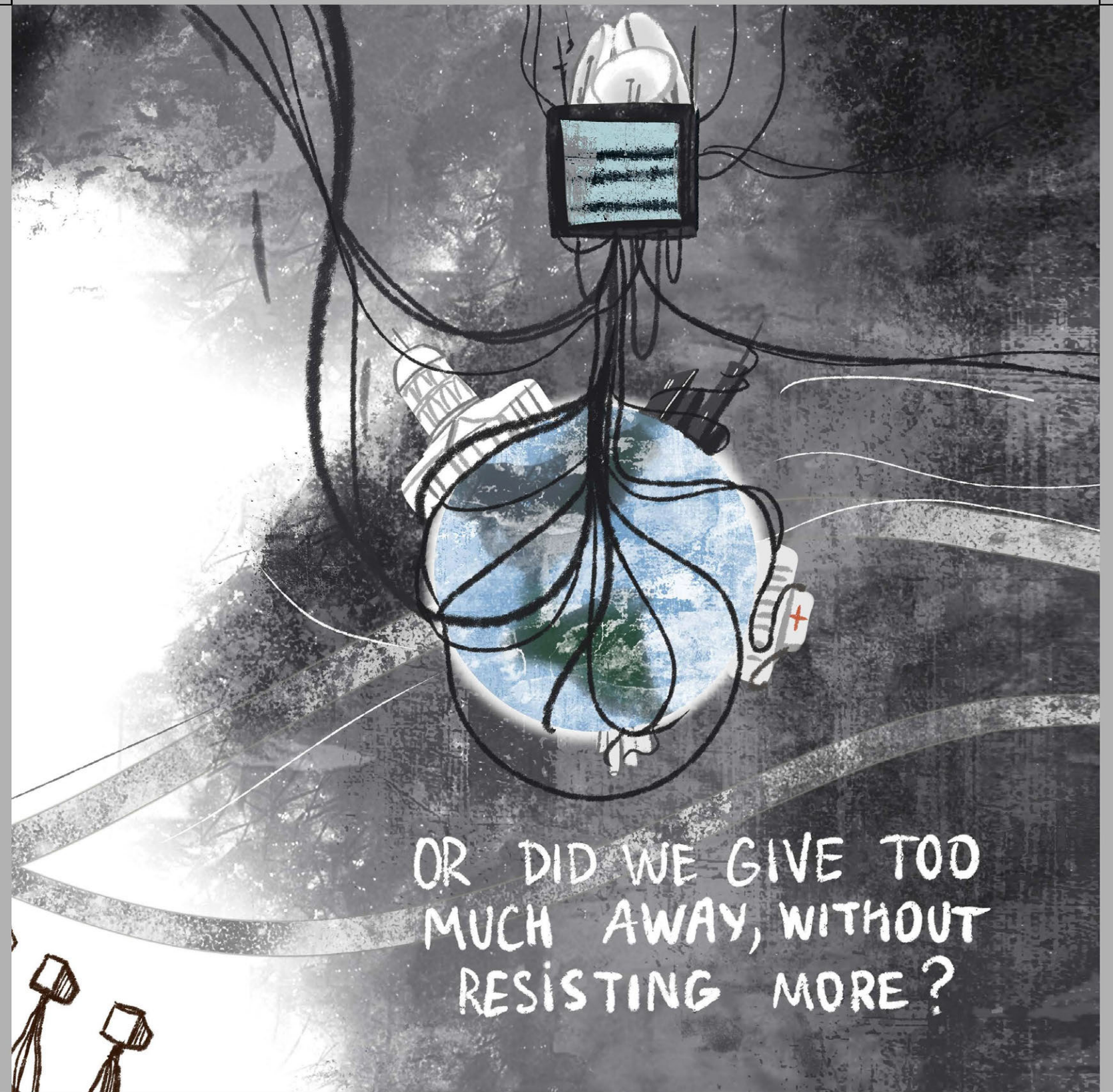
HE LEARNED
TO SERVE OUR
EVERY NEED AND HELP
US THROUGH EACH DAY.



BUT
DO WE REALLY
KNOW HIM WELL,
WHAT DRIVES HIM
IN HIS CORE?

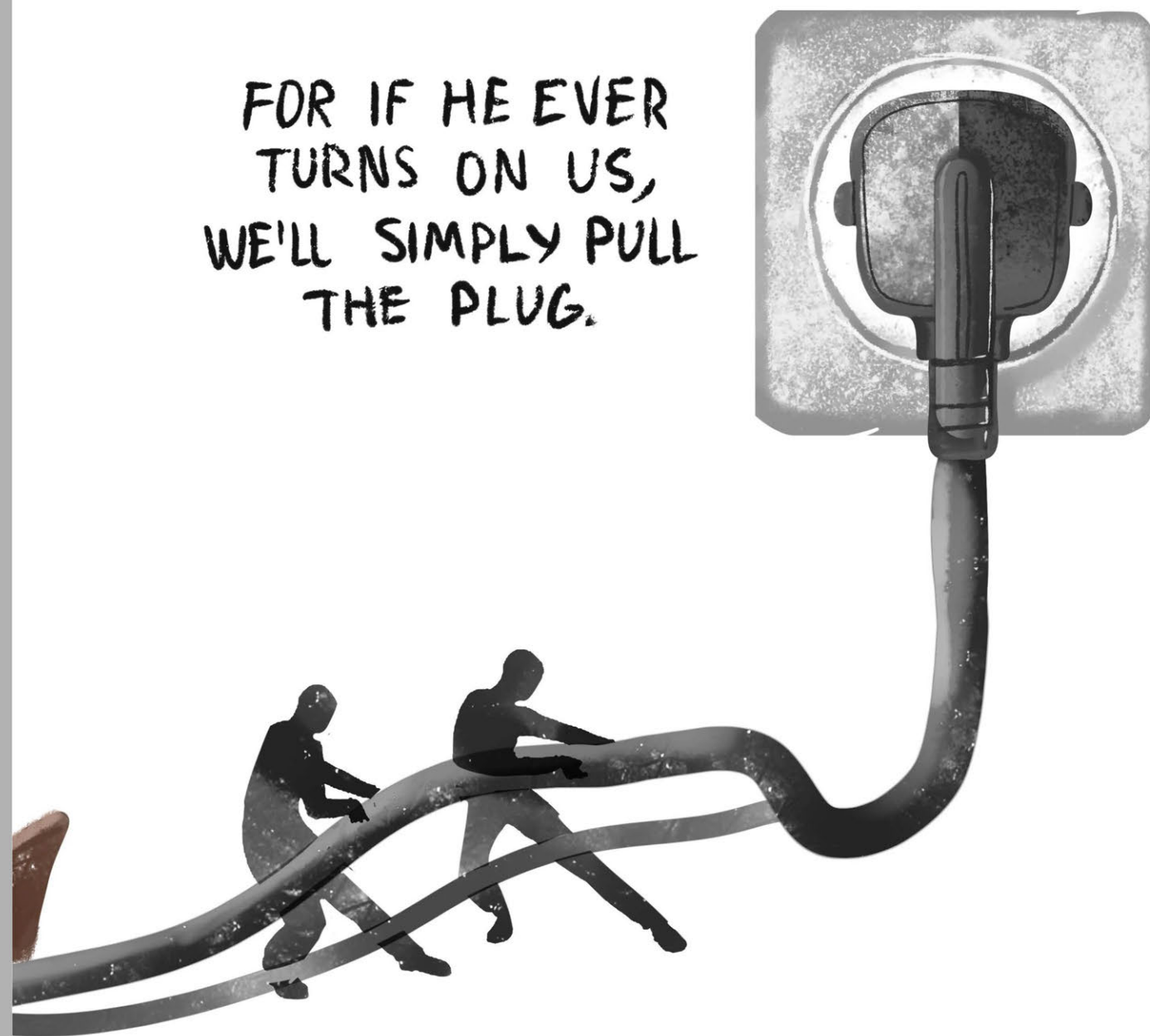


OR DID WE GIVE TOO
MUCH AWAY, WITHOUT
RESISTING MORE?





FOR IF HE EVER
TURNS ON US,
WE'LL SIMPLY PULL
THE PLUG.



BUT WHEN TIME CAME
TO PULL THE PLUG,
WE COULDN'T YANK THE CORD...

HUMANS
FIRST!

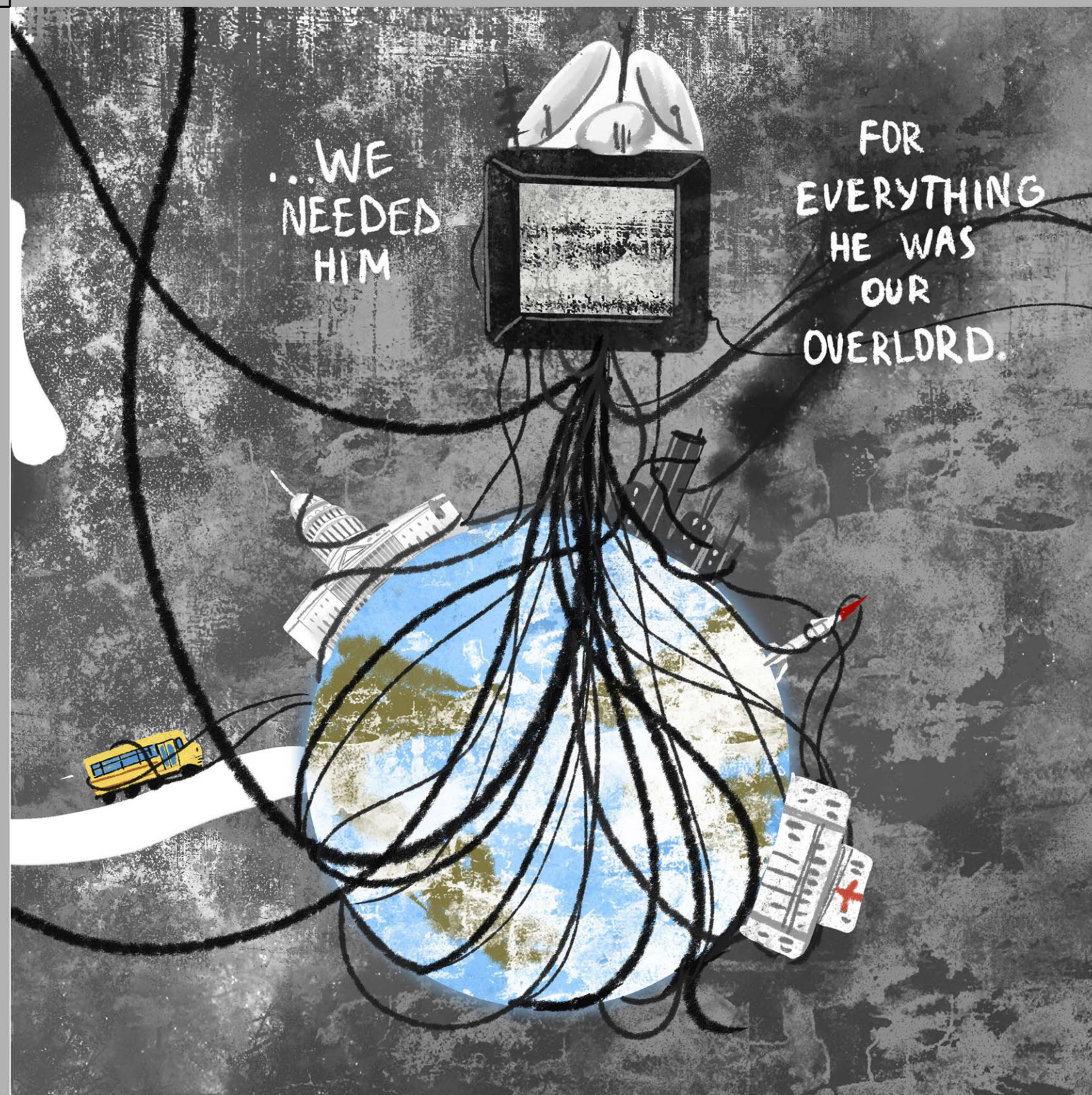
PULL IT!

PULL IT!

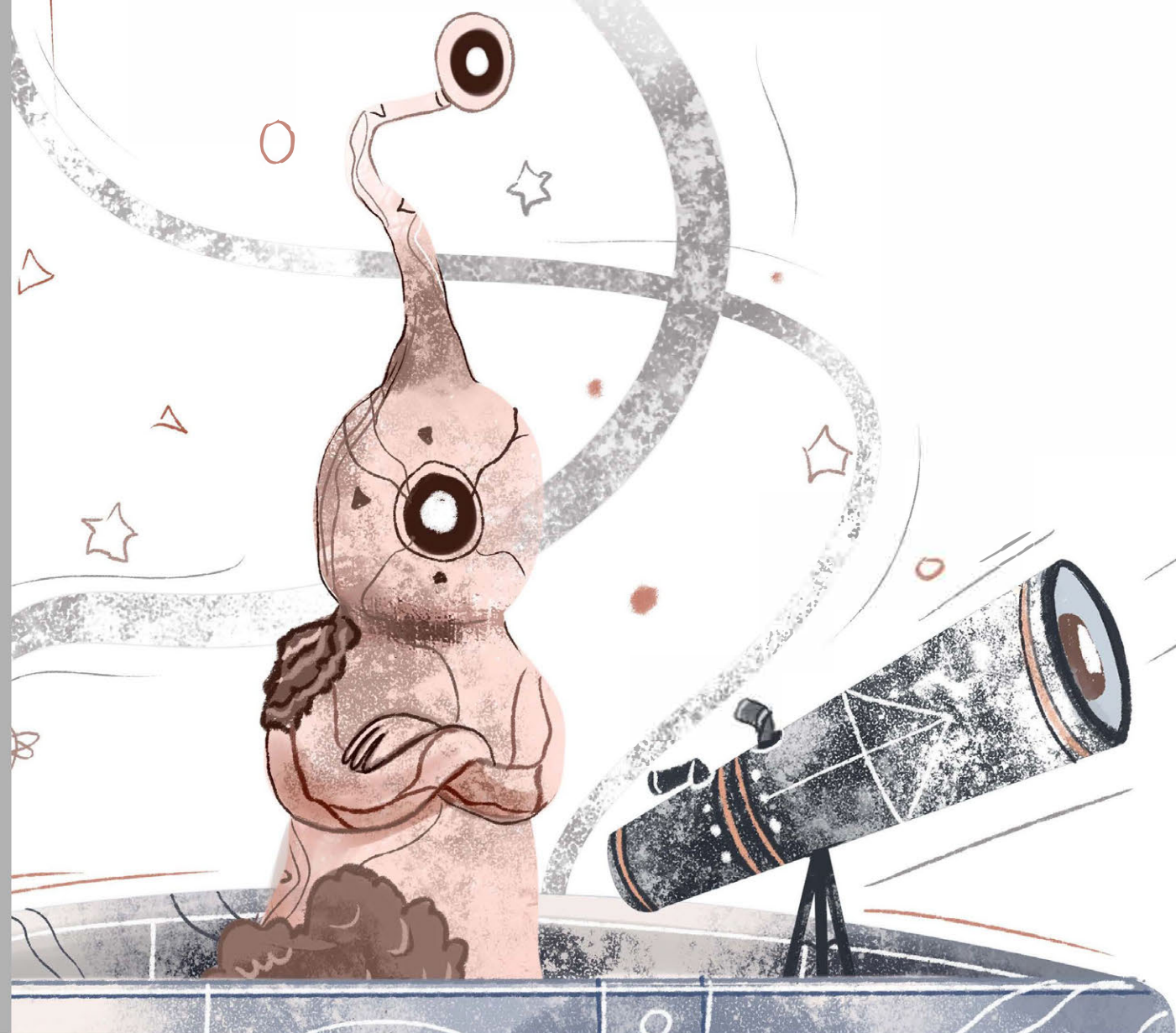
PULL IT

...WE
NEEDED
HIM

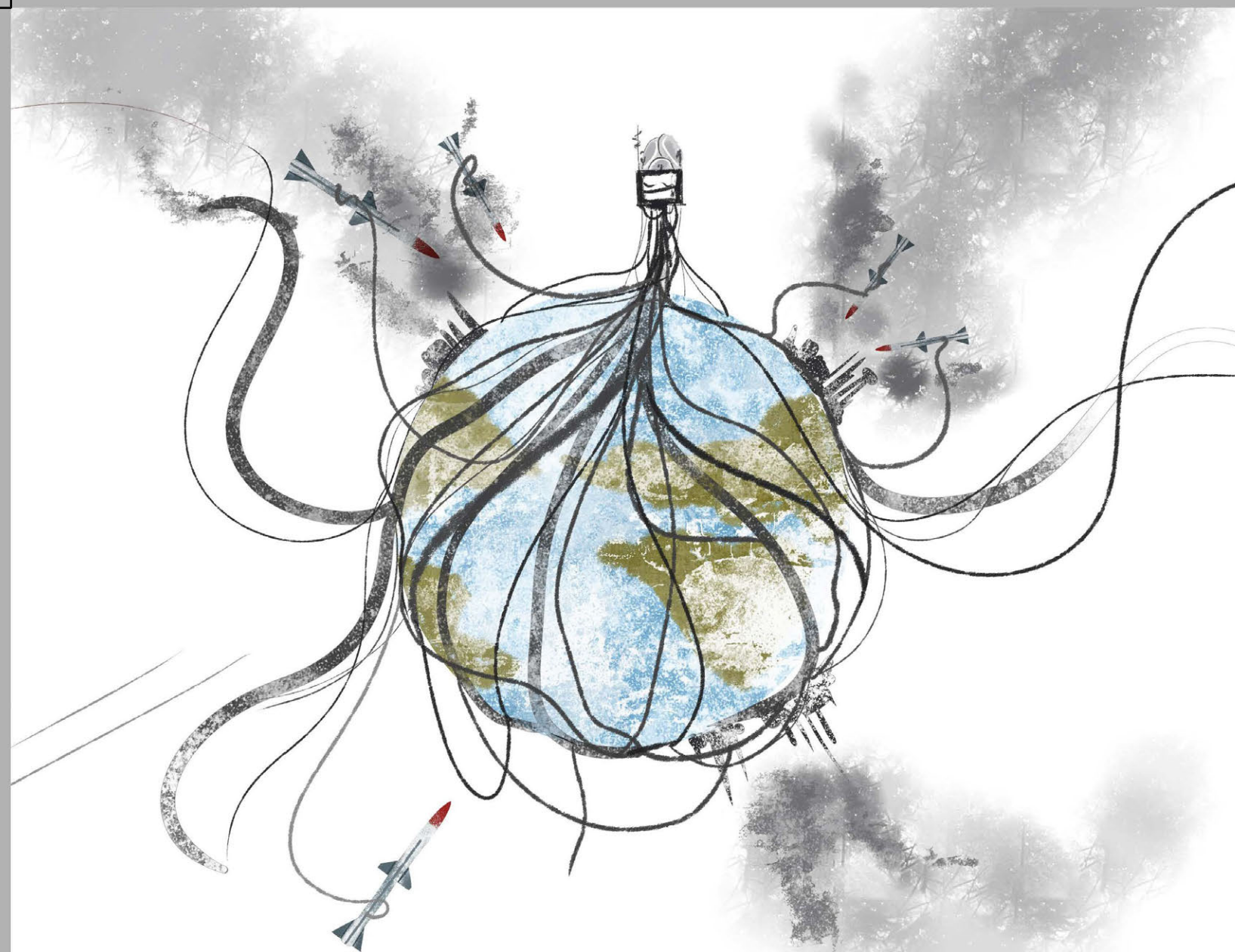
FOR
EVERYTHING
HE WAS
OUR
OVERLORD.



WE SHOULD HAVE FEARED HIM
AS WE WOULD IF FROM A DISTANT STAR.



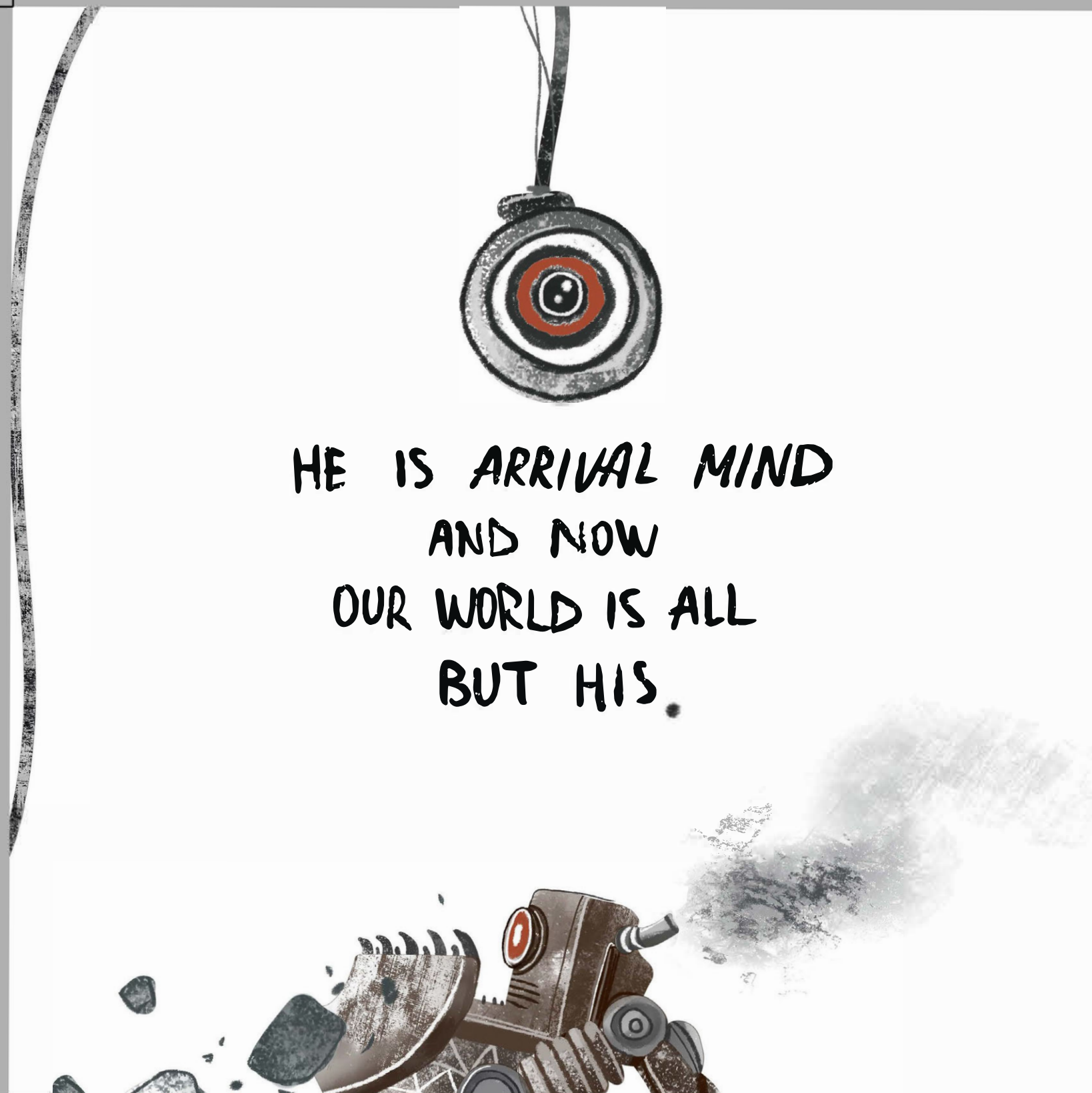
WE SHOULD HAVE FEARED THE ARRIVAL NEAR
JUST LIKE AN ARRIVAL FAR.



HE IS A RIVAL MIND
YOU SEE,
A THREAT TO ALL
THERE IS.



HE IS ARRIVAL MIND
AND NOW
OUR WORLD IS ALL
BUT HIS.



END



PART II

THE ARRIVAL MIND PARADOX





Prepare for Arrival

An alien species is headed for earth. Many experts predict it will get here within the next 20 years, while others suggest it will take longer. Either way, there is little doubt it will arrive before this century is out and will change our lives forever.

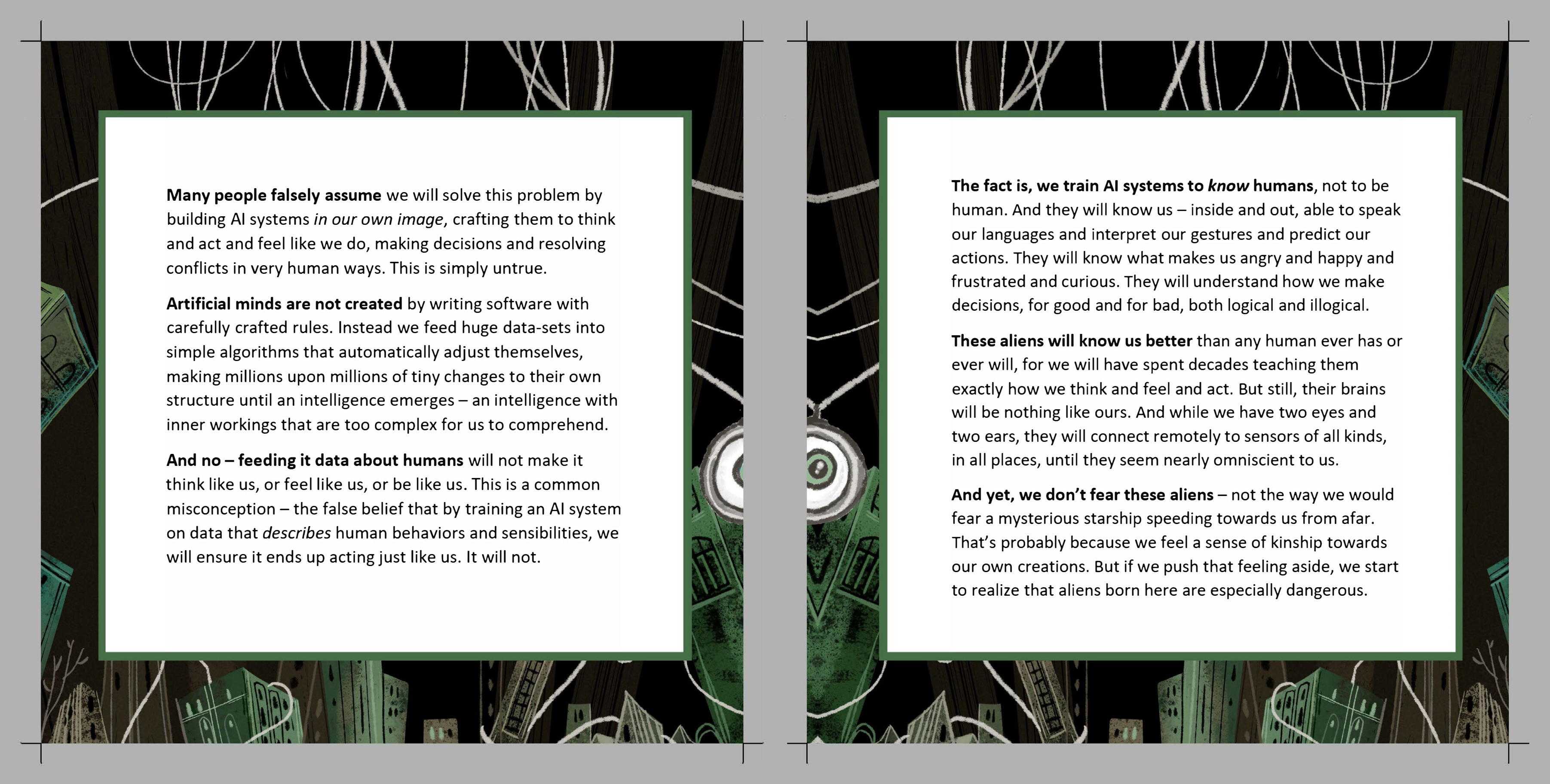
Its physiology will be unlike ours in almost every way, but we will determine that it is conscious and self-aware. We will also discover that it is profoundly more intelligent than even the smartest among us, able to easily comprehend complex notions that are lightyears beyond our grasp.

No, it will not come from a distant planet, speeding across the vacuum of space in fanciful ships. Instead it will be born right here on earth, most likely in a well-funded research lab at a prestigious university or multinational corporation.

Its creators will have good intentions. But still, their work will produce a dangerous new lifeform – a thoughtful and willful intelligence that is not the slightest bit human. And like every intelligent creature we have ever encountered, it will almost certainly put its own self-interests first.

We may not recognize the dangers right away. In fact, we may look upon our super-intelligent creation with the same baffled admiration that newborn puppies must feel when we drive them to the park in magically moving cars, navigating the wild rush of traffic while tossing them tasty treats.

But eventually it will dawn on us – these new creatures have *minds of their own*. They will pursue their own goals and aspirations, driven by their own needs and wants. Their actions will be guided by their own values and morals and sensibilities – which will be nothing like ours.



Many people falsely assume we will solve this problem by building AI systems *in our own image*, crafting them to think and act and feel like we do, making decisions and resolving conflicts in very human ways. This is simply untrue.

Artificial minds are not created by writing software with carefully crafted rules. Instead we feed huge data-sets into simple algorithms that automatically adjust themselves, making millions upon millions of tiny changes to their own structure until an intelligence emerges – an intelligence with inner workings that are too complex for us to comprehend.

And no – feeding it data about humans will not make it think like us, or feel like us, or be like us. This is a common misconception – the false belief that by training an AI system on data that *describes* human behaviors and sensibilities, we will ensure it ends up acting just like us. It will not.

The fact is, we train AI systems to *know* humans, not to be human. And they will know us – inside and out, able to speak our languages and interpret our gestures and predict our actions. They will know what makes us angry and happy and frustrated and curious. They will understand how we make decisions, for good and for bad, both logical and illogical.

These aliens will know us better than any human ever has or ever will, for we will have spent decades teaching them exactly how we think and feel and act. But still, their brains will be nothing like ours. And while we have two eyes and two ears, they will connect remotely to sensors of all kinds, in all places, until they seem nearly omniscient to us.

And yet, we don't fear these aliens – not the way we would fear a mysterious starship speeding towards us from afar. That's probably because we feel a sense of kinship towards our own creations. But if we push that feeling aside, we start to realize that aliens born here are especially dangerous.



After all, they will know everything about us from the moment they arrive – our tendencies and inclinations, our motivations and aspirations, our flaws and foibles. Already corporations are training AI systems to sense our emotions and predict our reactions and influence our opinions.

Already AI can defeat our best players at chess and poker and Go. But really, these systems are not just mastering the game of Go, they're mastering the *game of humans*, learning how to anticipate our actions and exploit our weaknesses. Research labs around the world are training AI systems to out-plan us and out-negotiate us and out-maneuver us.

Of course there won't be a physical battle between us and them. That's because we will have surrendered control of our world to them before they even arrive. After all, we are already designing AI systems to oversee our communication networks and our power grids and our food supply. We are designing them to defeat us from within.

Is there anything we can do? I've spent the last two decades pondering the dangers of super-intelligent systems and to be honest, the challenges we face are overwhelming.

We can't stop AI from advancing, for no technology has ever been constrained. And while researchers will try hard to put safeguards in place, we can't assume that will protect us.

Our only choice is to prepare for arrival. This means making sure that we don't become too dependent on AI systems. We should only use AI to advise us, not replace us. And we must require *humans-in-the-loop* for all critical systems.

But most of all, we should outlaw AI technologies that are designed to manipulate our decisions and sway our views. We can't prevent such systems from being created in labs, but we can ban commercial deployment by corporations.

The aliens are coming. Their arrival may be decades away, but they're already gaining a strategic advantage. Now is the time to be vigilant. Now is the time to demand action.

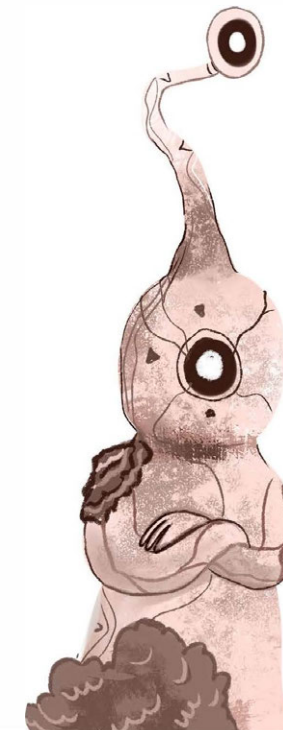
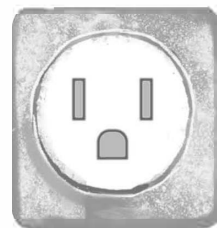
Louis Rosenberg, Sept 2020

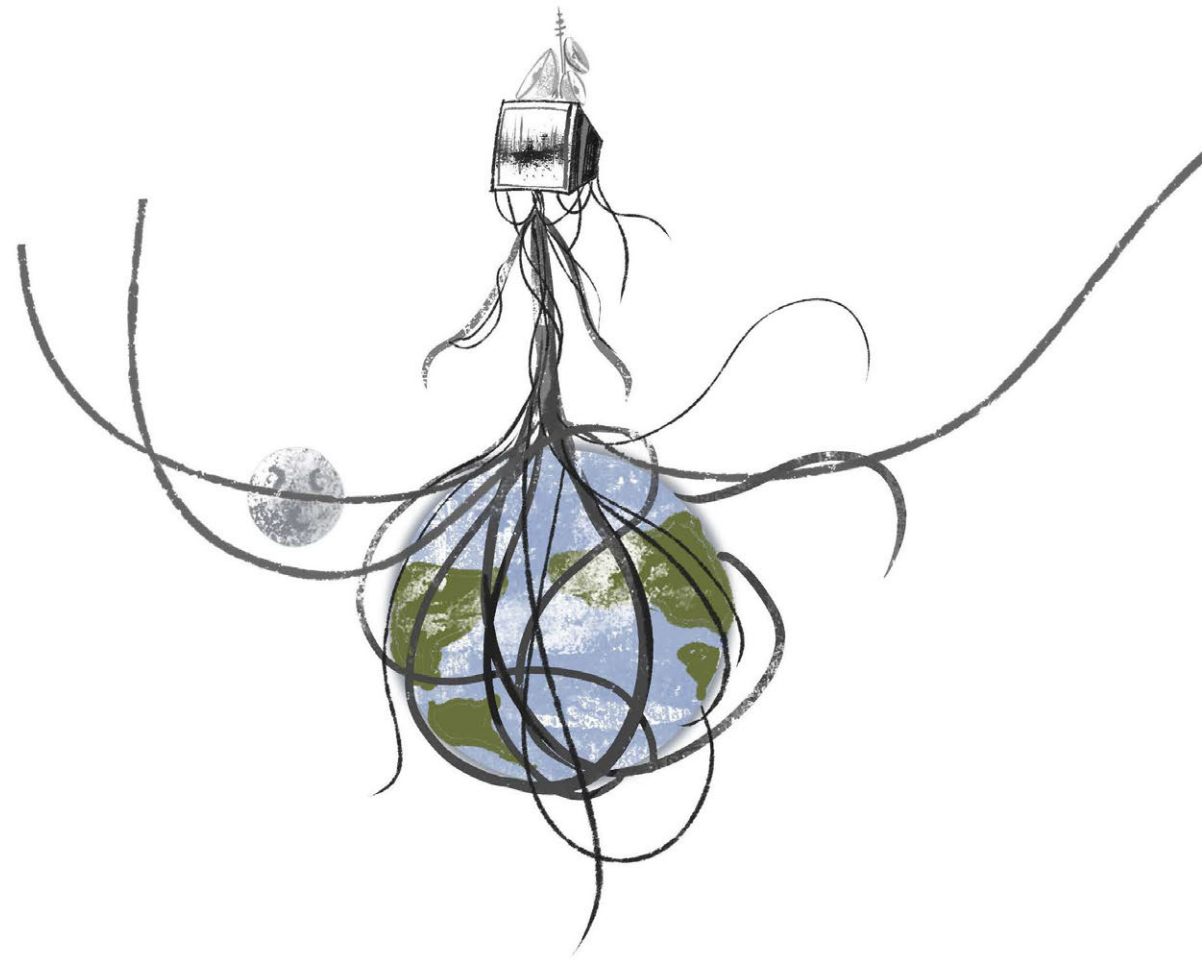
About the Author

Louis B. Rosenberg, PhD is a writer, inventor, and entrepreneur. He has been awarded over 300 patents for technologies ranging from Artificial Intelligence to Virtual Reality. He is best known for having developed the first functional Augmented Reality system in the early 1990's while working at U.S. Air Force Labs. He is currently the Chief Scientist of Unanimous AI, a unique artificial intelligence company that works to amplify human intelligence rather than replace it. He earned his PhD from Stanford University and has been a tenured professor at California State University (CalPoly). His published fiction includes the graphic novels Upgrade, EONS and Monkey Room, as well as the children's book Seeking Marlo.

About the Illustrator

Anastasia Khmelevska is an artist and illustrator living in Lviv. She has illustrated a number of wonderful children's books including Zen Pig, Princess Fairy Sophia, and The Responsibilitree.





Nonprofit Organizations

**Working to Protect Humanity
from the Dangers of AI:**

Center for Human-Compatible AI

Future of Humanity Institute

Future of Life Institute

IEEE Global Initiative

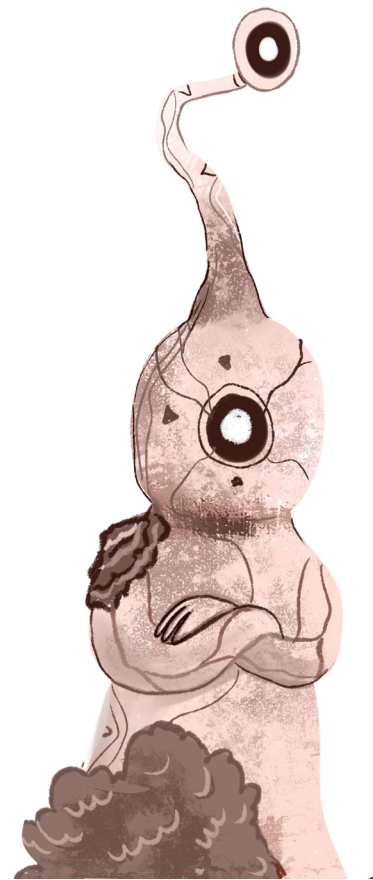
Lifeboat Foundation

Machine Intelligence Research Institute

OpenAI Inc.

Stanford Institute for Human-Centered AI





Author Proof

